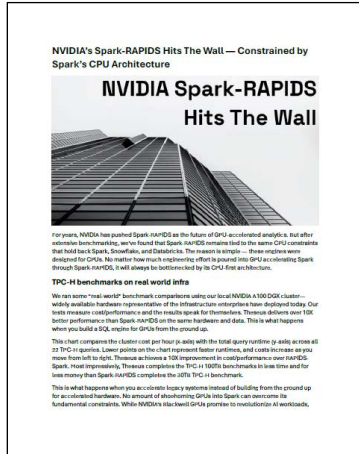




$b^3L5L!$ $\sigma \{ \mu \downarrow o \Xi w! tL5\{ I \uparrow \cup \sigma$
 $\text{£}^{\prime} \text{I} \text{S} \text{ } ^{\prime} \downarrow \text{PP} \quad / \text{Ü} \Omega \sigma \cup o \downarrow \Omega \text{S} \text{S} \text{ } \text{ö} \text{b} \text{I}$
 $\{ \mu \downarrow o \Xi \sigma / t \text{S} \text{ } ! o \text{ö}^{\prime} \text{I} \uparrow \cup \text{S} \text{S} \text{ö} \text{ö} \text{I} \text{ } \text{ö} \text{S}$

Thank you for downloading this Voltron Data resource. Carahsoft is the distributor for Voltron Data's AI and ML solutions available via NASA SEWP V, ITES-SW2, NASPO ValuePoint, and many more contract vehicles.

To learn how to take the next step toward acquiring Voltron Data's solutions, please check out the following resources and information:



For Voltron Data overview:
carah.io/Voltron-Data



For upcoming events:
carah.io/Voltron-Events



For additional resources:
carah.io/Voltron-Resources



For additional Artificial Intelligence:
carah.io/ai-solutions



To set up a meeting:
VoltronData@carahsoft.com



To purchase, check out the contract vehicles available for procurement:
carah.io/procurement

NVIDIA's Spark-RAPIDS Hits The Wall — Constrained by Spark's CPU Architecture



For years, NVIDIA has pushed Spark-RAPIDS as the future of GPU-accelerated analytics. But after extensive benchmarking, we've found that Spark-RAPIDS remains tied to the same CPU constraints that hold back Spark, Snowflake, and Databricks. The reason is simple — these engines were designed for CPUs. No matter how much engineering effort is poured into GPU accelerating Spark through Spark-RAPIDS, it will always be bottlenecked by its CPU-first architecture.

TPC-H benchmarks on real world infra

We ran some “real-world” benchmark comparisons using our local NVIDIA A100 DGX cluster—widely available hardware representative of the infrastructure enterprises have deployed today. Our tests measure cost/performance and the results speak for themselves. Theseus delivers over 10X better performance than Spark-RAPIDS on the same hardware and data. This is what happens when you build a SQL engine for GPUs from the ground up.

This chart compares the cluster cost per hour (x-axis) with the total query runtime (y-axis) across all 22 TPC-H queries. Lower points on the chart represent faster runtimes, and costs increase as you move from left to right. Theseus achieves a 10X improvement in cost/performance over RAPIDS-Spark. Most impressively, Theseus completes the TPC-H 100TB benchmarks in less time and for less money than Spark-RAPIDS completes the 30TB TPC-H benchmark.

This is what happens when you accelerate legacy systems instead of building from the ground up for accelerated hardware. No amount of shoehorning GPUs into Spark can overcome its fundamental constraints. While NVIDIA's Blackwell GPUs promise to revolutionize AI workloads,

they won't rewrite legacy CPU-based code to jump this wall. If NVIDIA wants true GPU-native data processing, it may need to rethink its entire approach—or just use Theseus and see what GPU acceleration looks like.

NEW TPC-DS benchmarks on the same infra

TPC-H and TPC-DS serve different purposes in benchmarking data analytics platforms. TPC-H is a strong measure of raw compute performance, where SQL optimizers will have limited impact on query runtime across the whole of the benchmark. TPC-DS, on the other hand, more aggressively stresses the query optimizer, and an engine's ability to simplify the complexity of a query plan for better performance. The Theseus SQL optimizer, while making great strides in the last few months due to hard work from our optimizer team, is still in its infancy when compared to tools with more mature optimizers, like Spark's. Even so, Theseus outperforms Spark and Spark-RAPIDS's better optimized query plans due to Theseus's raw performance as a GPU-native engine.

This is our first TPC-DS benchmark, and while we expect even better results over time, the initial performance is already impressive. Using the same cost-performance comparison as TPC-H, Theseus maintains an over 6X advantage over Spark-RAPIDS, further reinforcing the benefits of GPU-native query processing.

Benefits of an accelerator-native engine

Every new generation of GPU hardware delivers massive gains for software built to leverage accelerated infrastructure. There's no doubt that NVIDIA's Blackwell GPUs will be a game-changer for large-scale data analytics. With 10 TB/s chip-to-chip interconnects and a next-gen decompression engine delivering 900 GB/s of bidirectional NVLink bandwidth, Blackwell is designed to move massive datasets with unprecedented efficiency. At petabyte-scale, these hardware advancements are critical, enabling seamless integration of accelerated storage, networking, and compute into fully optimized, high-speed data pipelines.

But here's the problem: hardware alone isn't enough. While Spark-RAPIDS will see some gains from Blackwell's improvements, it remains fundamentally constrained by its CPU-first architecture. No amount of GPU retrofitting will fully unlock Blackwell's potential. Spark-RAPIDS is still bottlenecked by Spark's legacy architecture, preventing it from fully utilizing GPU memory, interconnects, and parallelism in the way Blackwell was engineered to support.

That's where Theseus has the edge. Unlike Spark-RAPIDS, which retrofits Spark to work with GPUs, Theseus was built from the ground up as an accelerator-native engine. We didn't merely adopt the CUDA-Verse — we were born in it, molded by it. Every component in Theseus is built to fully exploit NVIDIA's hardware and software ecosystem, from NVLink, UCX, GPUDirect Storage, NVComp, and libcudf (co-designed by our founders) and beyond.

Theseus is ready for Blackwell's new capabilities, from 900 GB/s interconnect speeds to its decompression engine, HBM bandwidth, and NVLink Switch. When GPUs are at the center of software design, optimizations happen at the lowest levels of the stack, where they matter most. Theseus isn't just another SQL engine that happens to run on GPUs. It's a SQL engine purpose-built to extract every ounce of performance from them.

The Cherry Pick

If you read our last benchmarking report you'll know we're not fans of using a simple column chart to compare two systems (see our [methodologies](#)). It typically supports cherry-picking of results that benefit the publishing vendor. Unfortunately, everyone is doing it and until people adopt what we think is a better way to benchmark with [the SPACE chart](#) we figure we would show how good we look when we cherry-pick.

Theseus outperforms Spark-RAPIDS by over 80X. It's starting to feel like a mic-drop moment.

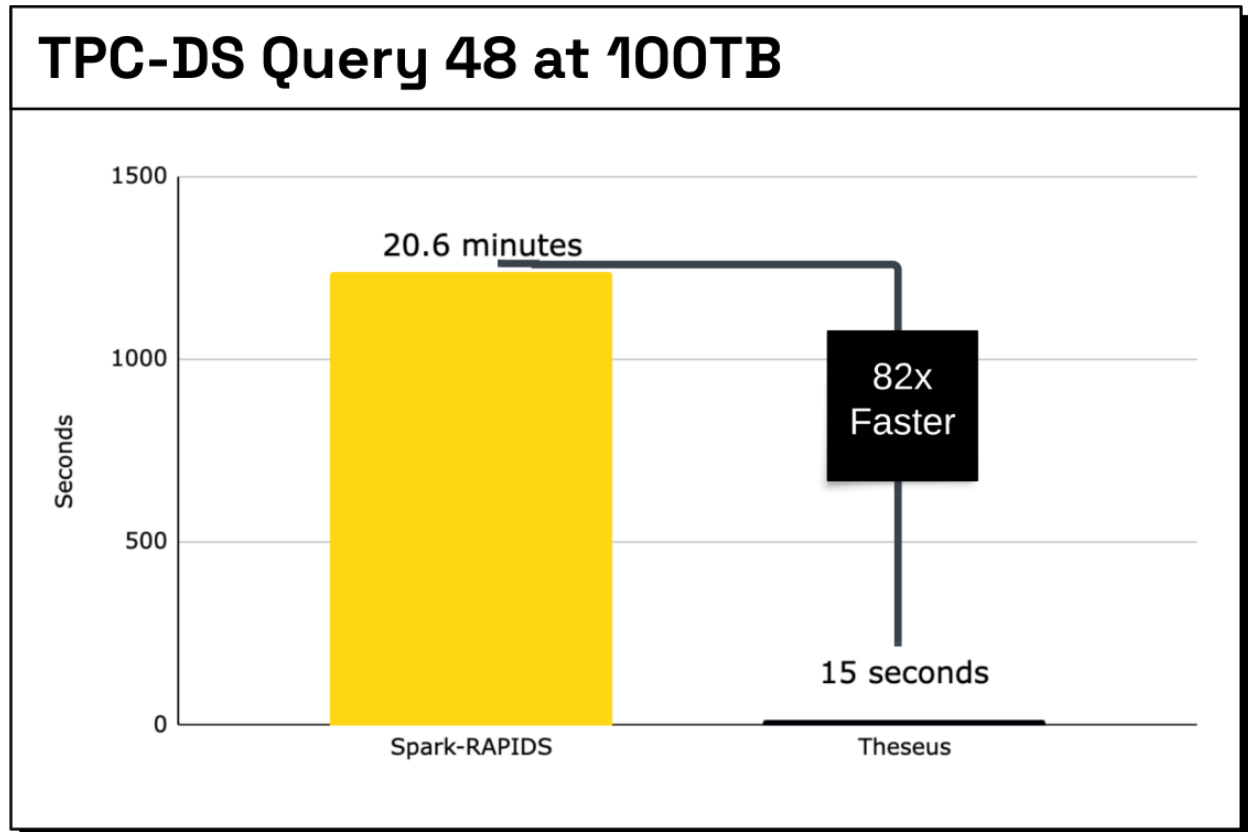


Figure 3. Performance Comparison of Voltron Data Theseus and Spark-RAPIDS running TPC-DS Query 48 at SF100000 using 6 nodes of an NVIDIA DGX A100 Cluster

Today, we unveiled new Theseus benchmark results for both TPC-H and TPC-DS, reaffirming its 10X cost-performance advantage over Spark-RAPIDS. But we're just getting started. We're now testing Theseus on NVIDIA's latest Blackwell B200 GPUs, and our engineers are eager to optimize for the GH200 Grace Hopper and GB200 Grace Blackwell Superchips to push performance even further.

We're also rolling out key updates to expand Theseus' capabilities:

- **NEW Multi-Engine Support** – We're adding DuckDB support as a CPU fallback via our composable data standards, with plans to evaluate other engines for smaller dataset price-performance. This broadens Theseus' applicability and will provide an easier onramp for customers.

- **NEW Multi-Silicon Support** – Customers are asking for more silicon options, so we’re bringing Theseus to ARM, Grace Hopper, Grace Blackwell, and MI300X architectures.
- **NEW Spark-Connect Integration** – Our engineers are working to enable Spark-Connect support, allowing customers to run legacy Spark code seamlessly while dramatically reducing onboarding friction and transition costs.

Wrap Up

The future of accelerated SQL analytics is here. Stay tuned as we continue to push the boundaries of performance, efficiency, and cost savings.

For more details on our benchmarking process, see our [Theseus Benchmarking Report](#), where we break down the [results](#), [methodologies](#), [FAQs](#), and [system specifications](#).

Remember, it’s all about how fast your queries run and what it costs to run them. Theseus delivers. Join our mission to slash the cost of analytics and give [Theseus a try](#).